

Inference Scaling for Long-Context Retrieval Augmented Generation

Zhenrui Yue^{*1,2}, Honglei Zhuang^{*1}, Aijun Bai¹, Kai Hui¹, Rolf Jagerman¹, Hansi Zeng^{1,3}, Zhen Qin¹, Dong Wang², Xuanhui Wang¹, Michael Bendersky¹

¹Google DeepMind

²University of Illinois Urbana-Champaign

³University of Massachusetts Amherst

^{*}Equal contribution



<https://arxiv.org/abs/2410.04343>

Introduction



Introduction

Inference Scaling

A series of recent studies show that increasing the amount of inference-time computation can be similar (Agarwal et al., 2024), if not more effective (Snell et al., 2024) than allocating those computation to training in some scenarios.

Examples include:

- Scaling the number of examples in ICL
- Scaling best-of-N samples along with sequential revisions
- Scaling reasoning iterations (e.g., OpenAI's o1 model)
- ...

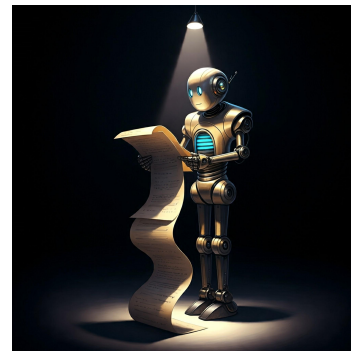
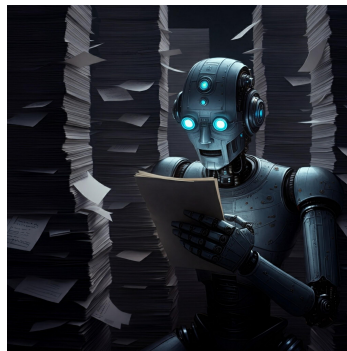


Introduction

Inference Scaling for RAG

With the advances in long-context LLMs, recent studies also attempt to better leverage the full context-length in retrieval-augmented generation (RAG) tasks.

Studies on inference scaling for RAG mostly focus on scaling the **number of documents** (Ram et al., 2023; Shao et al., 2024; Lee et al., 2024; Xu et al., 2024;) or the **length of documents** (Jiang et al., 2024).



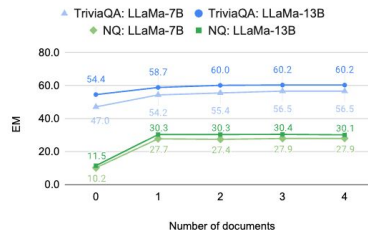
Introduction

Inference Scaling for RAG

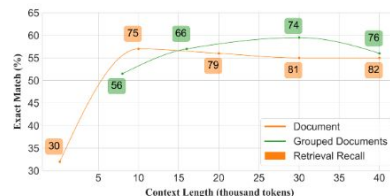
With the advances in long-context LLMs, recent studies also attempt to better leverage the full context-length in retrieval-augmented generation (RAG) tasks.

Studies on inference scaling for RAG mostly focus on scaling the **number of documents** (Ram et al., 2023; Shao et al., 2024; Lee et al., 2024; Xu et al., 2024;) or the **length of documents** (Jiang et al., 2024).

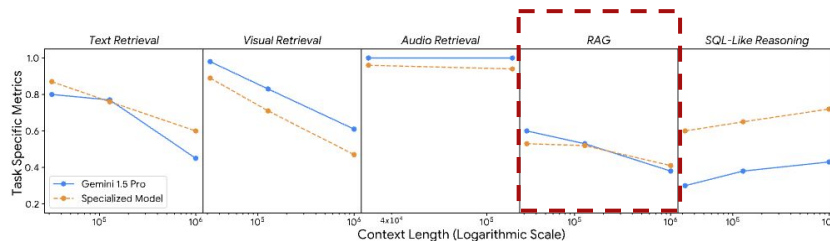
... And RAG performance does not always increase as the retrieved context increases!



(Ram et al. 2023)



(Jiang et al. 2024)



(Lee et al. 2024)

Introduction

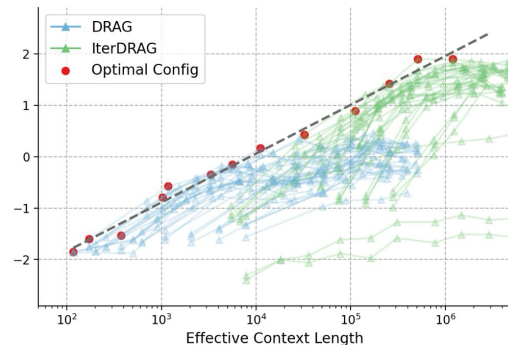
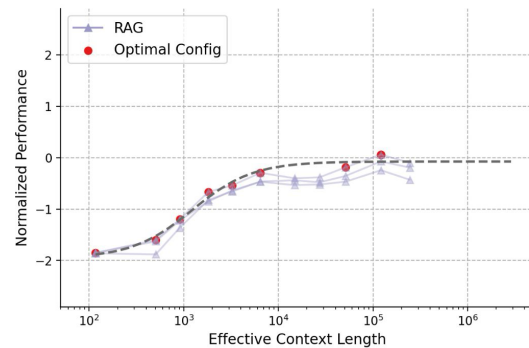
More Comprehensive Inference Scaling for RAG

In this work, we conduct a more comprehensive study on inference scaling for RAG.

We study two more strategies to leverage inference computation in RAG tasks:

1. **Demonstration-Based RAG (DRAG):** Adding RAG demonstrations as in-context examples.
2. **Iterative Demonstration-Based RAG (IterDRAG):** Iteratively apply demonstration-based RAG to solve more challenging, multi-hop queries.

And we show that when optimally configured, these strategies enable RAG performance to increase (almost) linearly with the order of magnitude of the amount of inference computation.



Inference Scaling Strategies

2

Inference Scaling Strategies

Definition of “Inference Computation”

In our study, we measure the **inference computation** by “**effective context length**”:

- For *single-round* strategies, this is equivalent to the input context length to the LLM
- For *multi-round* strategies, this is the sum of the input context lengths for every rounds of LLM calls

This aligns well with the pricing model of many commercial LLMs, as the output token number for our tasks is often limited.

We consider a fixed-budget setting, where users are given a fixed budget of **effective context length** L_{max}

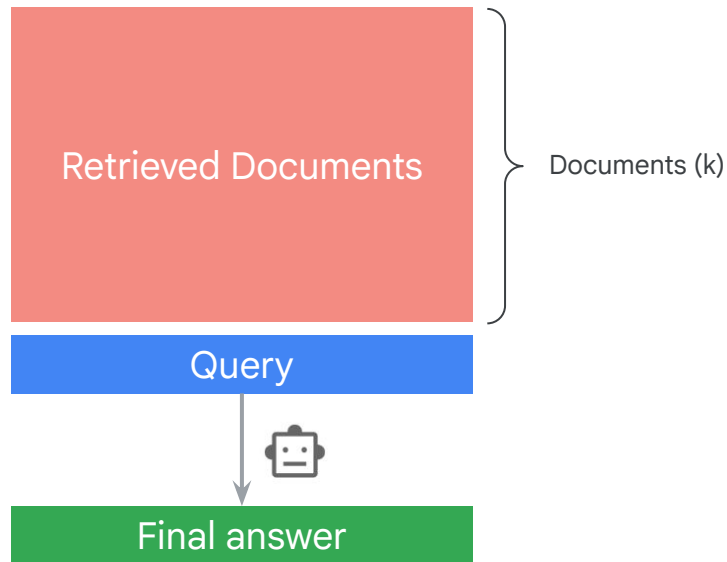
Inference Scaling Strategies

Vanilla RAG

In the vanilla RAG strategy, a retriever will retrieve k documents based on the query. The retrieved documents and the query are provided to the LLM as input, and the LLM outputs the answer.

To fully leverage the context window of LLMs, we can adjust the following parameter:

- Number of documents k



Inference Scaling Strategies

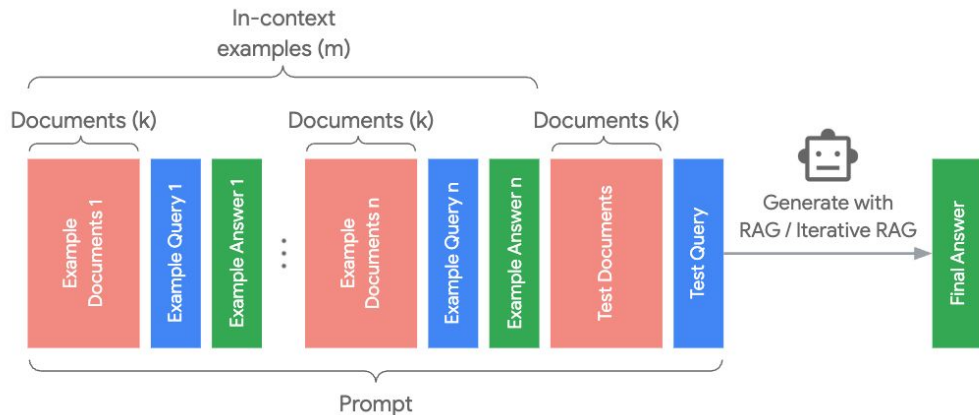
Demonstration-Based RAG (DRAG)

Adding demonstrations as in-context examples, where each demonstration include a complete RAG call: retrieved documents, query and answer.

Ideally, demonstrations allow models to learn how to locate the most relevant information and follow the formatting convention of answers.

This strategy's effective context length can be controlled by 2 parameters

1. Number of documents k
2. Number of in-context examples m



Inference Scaling Strategies

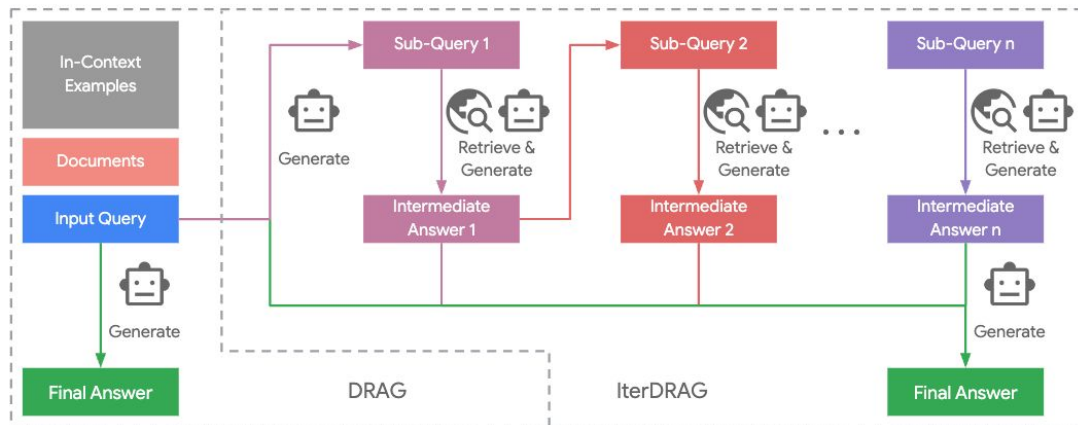
Iterative Demonstration-Based RAG (IterDRAG)

Based on the DRAG strategy, we can further allow the model to iteratively issue sub-queries based on the answer from previous rounds.

Ideally, the iterative process allows models to tackle queries requiring multi-hop reasoning.

This strategy's effective context length can be controlled by 3 parameters

1. Number of documents k
2. Number of in-context examples m
3. Number of iterations n



RAG Performance and Inference Computation Scale

3

RAG Performance and Inference Computation Scale

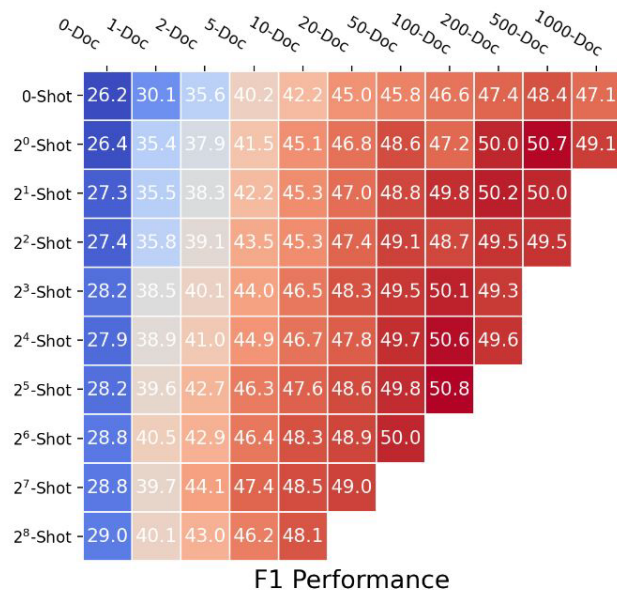
Fixed Budget Optimal Performance

Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.



RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

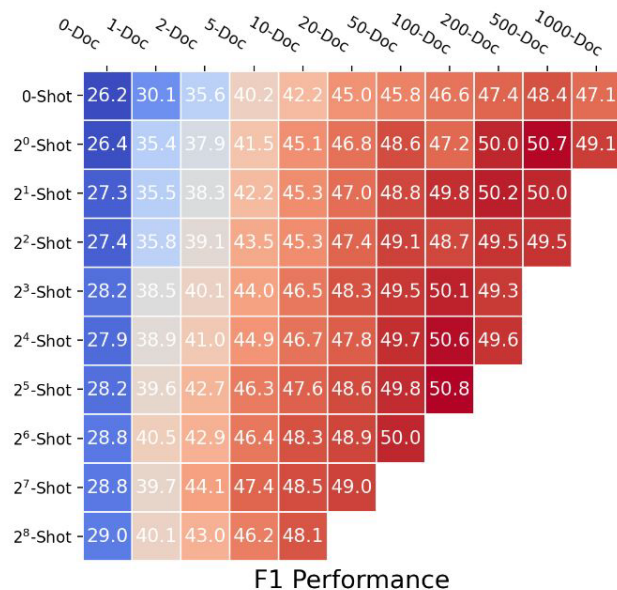
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

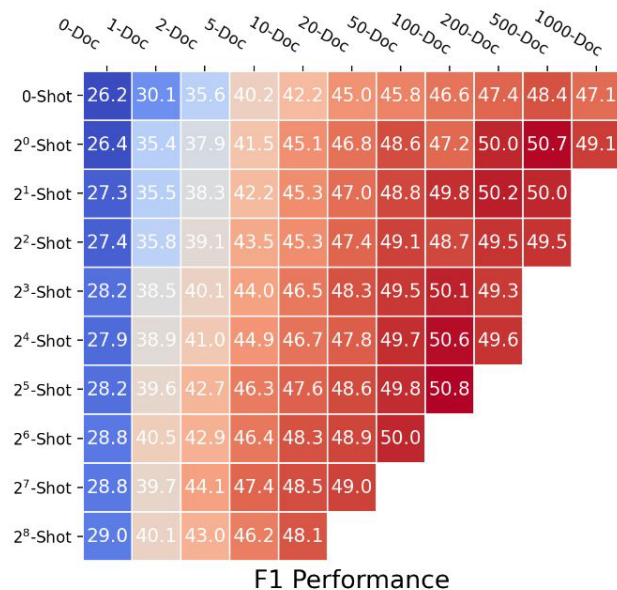
For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

Assuming we can always find the optimal configuration, how will the performance scale with the budget?





RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

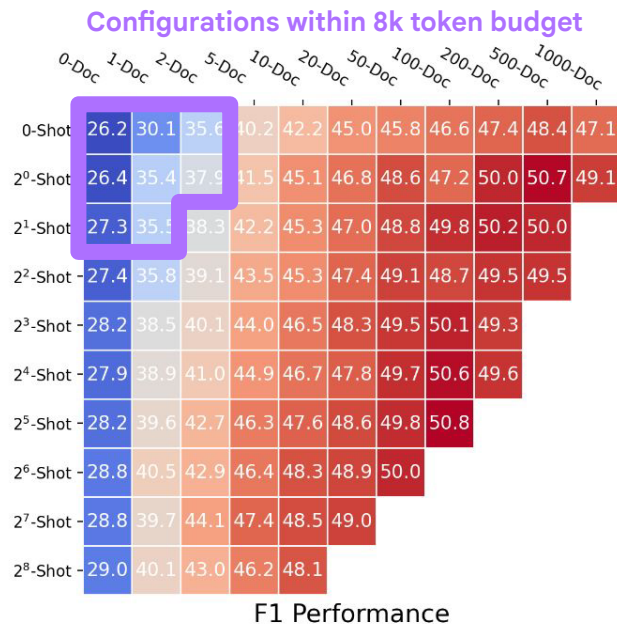
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

Among all these configurations, we can find the optimal RAG performance, denoted as $P^*(L_{max})$



RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

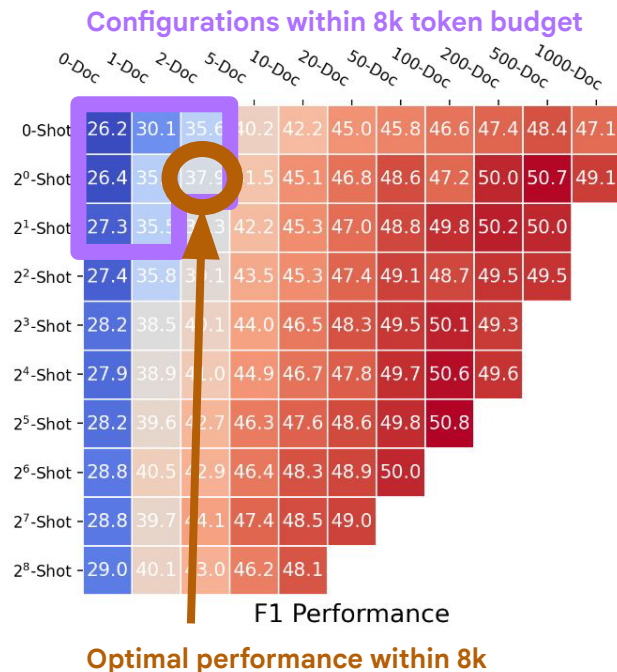
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

Among all these configurations, we can find the optimal RAG performance, denoted as $P^*(L_{max})$



RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

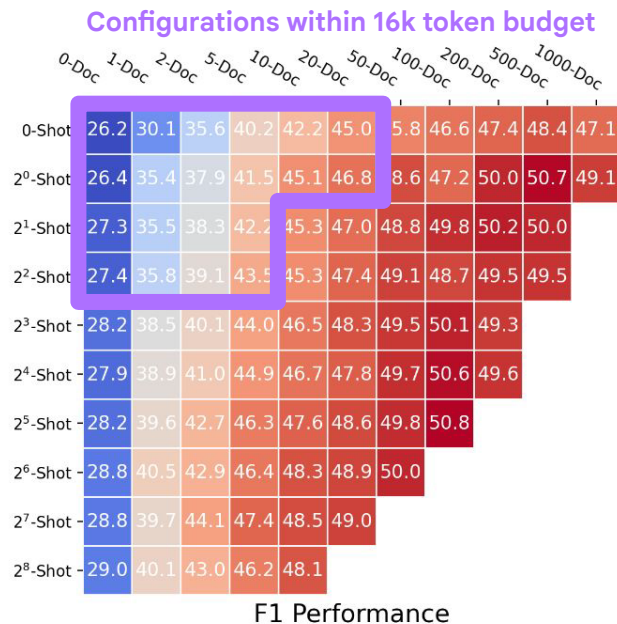
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

Among all these configurations, we can find the optimal RAG performance, denoted as $P^*(L_{max})$



RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

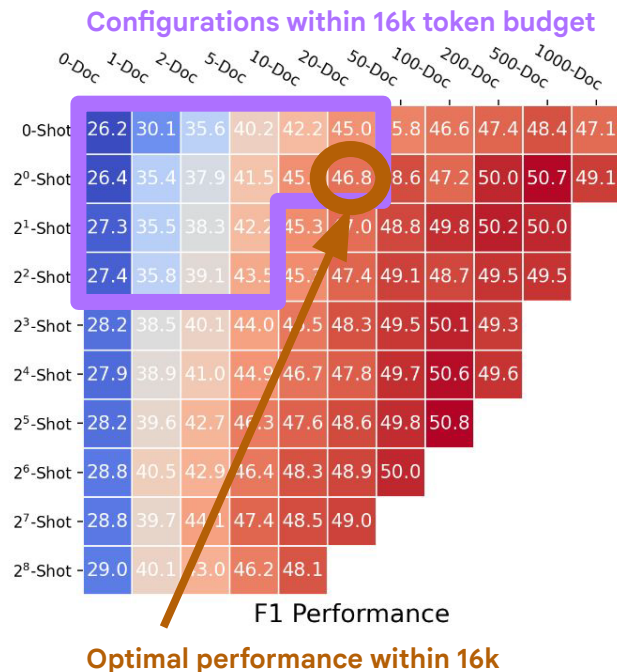
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

Among all these configurations, we can find the optimal RAG performance, denoted as $P^*(L_{max})$



RAG Performance and Inference Computation Scale

Fixed Budget Optimal Performance

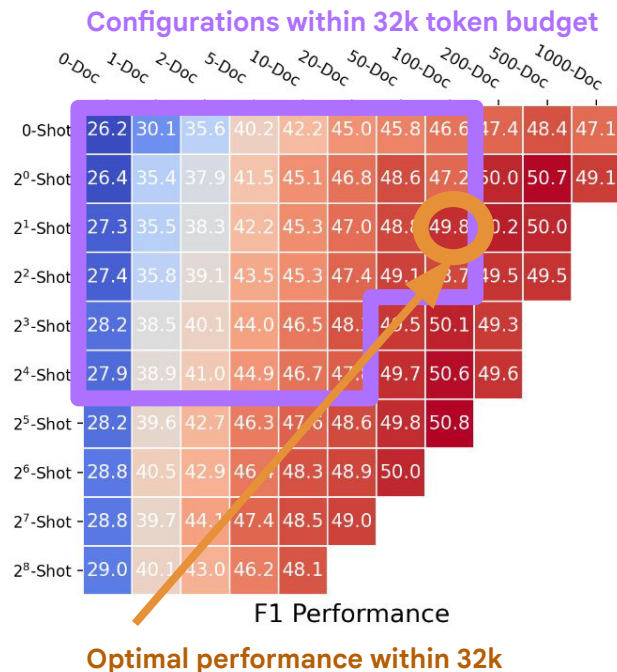
Given a fixed budget L_{max} and a strategy (DRAG, IterDRAG etc.), there can be different configurations satisfying the budget

For example, if the budget is 8k tokens, and the strategy is DRAG, the following configurations can all satisfy the budget:

- Number of documents $k=20$, Number of demos $m=1$
- Number of documents $k=10$, Number of demos $m=2$
- Number of documents $k=5$, Number of demos $m=4$
- ...

We enumerate a set of configuration combination for each strategy.

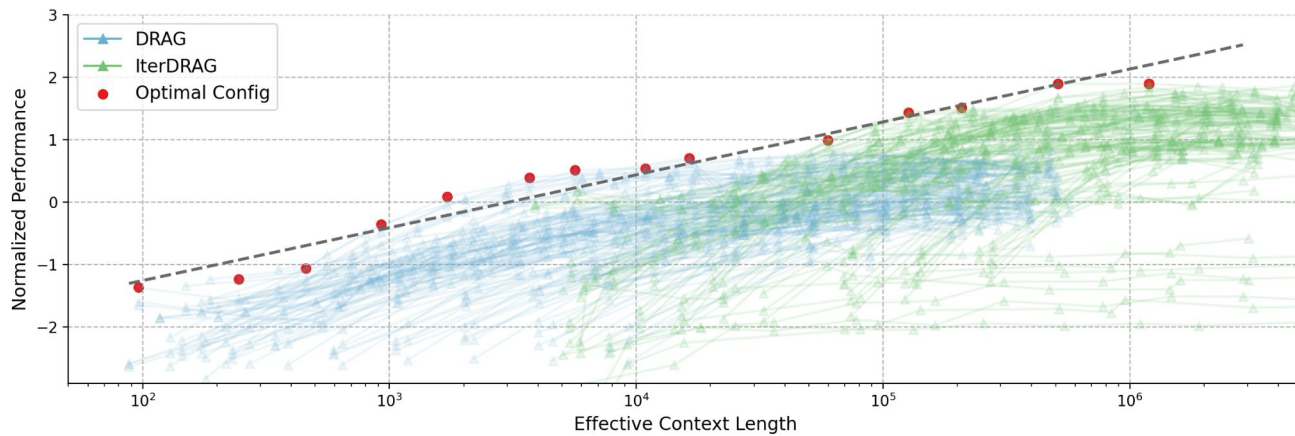
Among all these configurations, we can find the optimal RAG performance, denoted as $P^*(L_{max})$



RAG Performance and Inference Computation Scale

RAG Performance vs. Inference Computation Scale

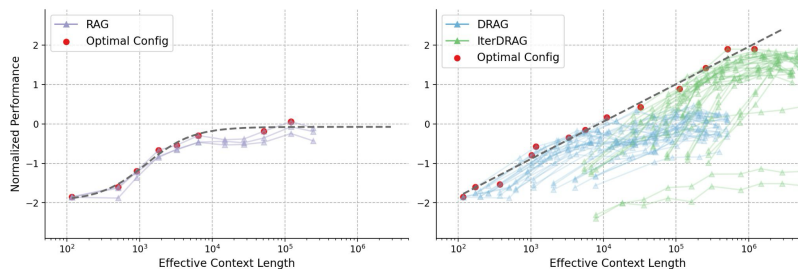
Plotting $P^*(L_{max})$ with L_{max} for both strategies on all datasets with three metrics (EM, F1, Accuracy) and normalize them.



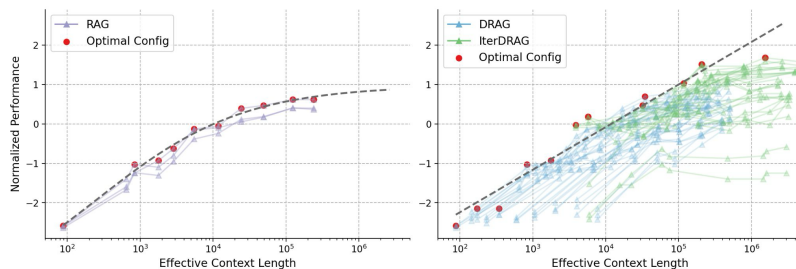
RAG Performance and Inference Computation Scale

RAG Performance vs. Inference Computation Scale

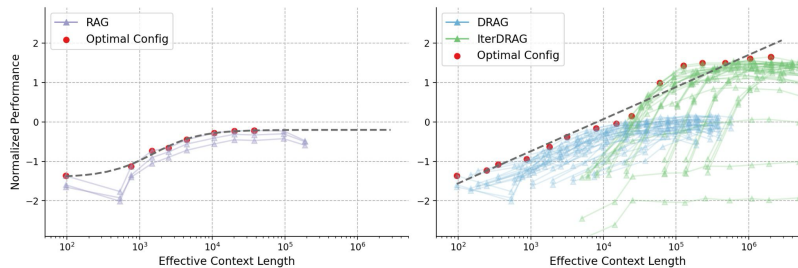
Comparing to only scaling the number of document in vanilla RAG on different datasets.



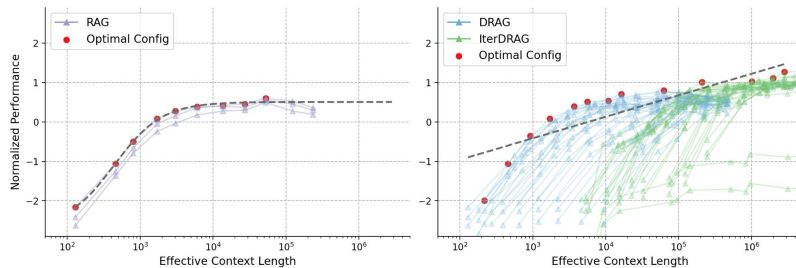
MuSiQue



Bamboogle



2WikiMultiHopQA



HotpotQA

RAG Performance and Inference Computation Scale

Comparing Strategies

Evaluated on 4 multi-hop open-book question answering datasets.

Baselines:

- **Zero-Shot QA (ZS QA):** 0 retrieved document, 0 demonstration
- **Many-Shot QA (MS QA):** 0 retrieved document, many demonstrations
- **RAG:** many retrieved documents, 0 demonstration

Optimal performance of different methods with varying maximum effective context lengths.

$L_{\max}$	Method	Bambooole			HotpotQA			MuSiQue			2WikiMultiHopQA		
		EM	F1	Acc	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc
16k	ZS QA	16.8	25.9	19.2	22.7	32.0	25.2	5.0	13.2	6.6	28.3	33.5	30.7
	MS QA	24.0	30.7	24.8	24.6	34.0	26.2	7.4	16.4	8.5	33.2	37.5	34.3
	RAG	44.0	54.5	45.6	44.2	57.9	49.2	12.3	21.5	15.3	42.3	49.3	46.5
	DRAG	44.0	55.2	45.6	45.5	58.5	50.2	14.5	24.6	16.9	45.2	53.5	50.5
	IterDRAG	46.4	56.2	51.2	36.0	47.4	44.4	8.1	17.5	12.2	33.2	38.8	43.8
32k	RAG	48.8	56.2	49.6	44.2	58.2	49.3	12.3	21.5	15.3	42.9	50.6	48.0
	DRAG	48.8	59.2	50.4	46.9	60.3	52.0	15.4	26.0	17.3	45.9	53.7	51.4
	IterDRAG	46.4	56.2	52.0	38.3	49.8	44.4	12.5	23.1	19.7	44.3	54.6	56.8
128k	RAG	51.2	60.3	52.8	45.7	59.6	50.9	14.0	23.7	16.8	43.1	50.7	48.4
	DRAG	52.8	62.3	54.4	47.4	61.3	52.2	15.4	26.0	17.9	47.5	55.3	53.1
	IterDRAG	63.2	74.8	68.8	44.8	59.4	52.8	17.3	28.0	24.5	62.3	73.8	74.6
1M	DRAG	56.0	62.9	57.6	47.4	61.3	52.2	15.9	26.0	18.2	48.2	55.7	53.3
	IterDRAG	65.6	75.6	68.8	48.7	63.3	55.3	22.2	34.3	30.5	65.7	75.2	76.4
5M	IterDRAG	65.6	75.6	68.8	51.7	64.4	56.4	22.5	35.0	30.5	67.0	75.2	76.9

RAG Performance and Inference Computation Scale

Comparing Strategies

Evaluated on 4 multi-hop open-book question answering datasets.

Baselines:

- **Zero-Shot QA (ZS QA):** 0 retrieved document, 0 demonstration
- **Many-Shot QA (MS QA):** 0 retrieved document, many demonstrations
- **RAG:** many retrieved documents, 0 demonstration

Optimal performance of different methods with varying maximum effective context lengths.

$L_{\max}$	Method	Bamboogle			HotpotQA			MuSiQue			2WikiMultiHopQA		
		EM	F1	Acc	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc
16k	ZS QA	16.8	25.9	19.2	22.7	32.0	25.2	5.0	13.2	6.6	28.3	33.5	30.7
	MS QA	24.0	30.7	24.8	24.6	34.0	26.2	7.4	16.4	8.5	33.2	37.5	34.3
	RAG	44.0	54.5	45.6	44.2	57.9	49.2	12.3	21.5	15.3	42.3	49.3	46.5
	DRAG IterDRAG	44.0 46.4	55.2 56.2	45.6 51.2	45.5 36.0	58.5 47.4	50.2 44.4	14.5 8.1	24.6 17.5	16.9 12.2	45.2 33.2	53.5 38.8	50.5 43.8
32k	RAG	48.8	56.2	49.6	44.2	58.2	49.3	12.3	21.5	15.3	42.9	50.6	48.0
	DRAG IterDRAG	48.8 46.4	59.2 56.2	50.4 52.0	46.9 38.3	60.3 49.8	52.0 44.4	15.4 12.5	26.0 23.1	17.3 19.7	45.9 44.3	53.7 54.6	51.4 56.8
	RAG	51.2	60.3	52.8	45.7	59.6	50.9	14.0	23.7	16.8	43.1	50.7	48.4
128k	DRAG IterDRAG	52.8 63.2	62.3 74.8	54.4 68.8	47.4 44.8	61.3 59.4	52.2 52.8	15.4 17.3	26.0 28.0	17.9 24.5	47.5 62.3	55.3 73.8	53.1 74.6
1M	DRAG IterDRAG	56.0 65.6	62.9 75.6	57.6 68.8	47.4 48.7	61.3 63.3	52.2 55.3	15.9 22.2	26.0 34.3	18.2 30.5	48.2 65.7	55.7 75.2	53.3 76.4
5M	IterDRAG	65.6	75.6	68.8	51.7	64.4	56.4	22.5	35.0	30.5	67.0	75.2	76.9

DRAG and IterDRAG consistently achieve higher performance than other strategies on different $L_{\max}$

RAG Performance and Inference Computation Scale

Comparing Strategies

Evaluated on 4 multi-hop open-book question answering datasets.

Baselines:

- **Zero-Shot QA (ZS QA):** 0 retrieved document, 0 demonstration
- **Many-Shot QA (MS QA):** 0 retrieved document, many demonstrations
- **RAG:** many retrieved documents, 0 demonstration

Optimal performance of different methods with varying maximum effective context lengths.

$L_{\max}$	Method	Bamboogle			HotpotQA			MuSiQue			2WikiMultiHopQA		
		EM	F1	Acc	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc
16k	ZS QA	16.8	25.9	19.2	22.7	32.0	25.2	5.0	13.2	6.6	28.3	33.5	30.7
	MS QA	24.0	30.7	24.8	24.6	34.0	26.2	7.4	16.4	8.5	33.2	37.5	34.3
	RAG	44.0	54.5	45.6	44.2	57.9	49.2	12.3	21.5	15.3	42.3	49.3	46.5
	DRAG	44.0	55.2	45.6	45.5	58.5	50.2	14.5	24.6	16.9	45.2	53.5	50.5
	IterDRAG	46.4	56.2	51.2	36.0	47.4	44.4	8.1	17.5	12.2	33.2	38.8	43.8
32k	RAG	48.8	56.2	49.6	44.2	58.2	49.3	12.3	21.5	15.3	42.9	50.6	48.0
	DRAG	48.8	59.2	50.4	46.9	60.3	52.0	15.4	26.0	17.3	45.9	53.7	51.4
	IterDRAG	46.4	56.2	52.0	38.3	49.8	44.4	12.5	23.1	19.7	44.3	54.6	56.8
128k	RAG	51.2	60.3	52.8	45.7	59.6	50.9	14.0	23.7	16.8	43.1	50.7	48.4
	DRAG	52.8	62.3	54.4	47.4	61.3	52.2	15.4	26.0	17.9	47.5	55.3	53.1
	IterDRAG	63.2	74.8	68.8	44.8	59.4	52.8	17.3	28.0	24.5	62.3	73.8	74.6
1M	DRAG	56.0	62.9	57.6	47.4	61.3	52.2	15.9	26.0	18.2	48.2	55.7	53.3
	IterDRAG	65.6	75.6	68.8	48.7	63.3	55.3	22.2	34.3	30.5	65.7	75.2	76.4
5M	IterDRAG	65.6	75.6	68.8	51.7	64.4	56.4	22.5	35.0	30.5	67.0	75.2	76.9

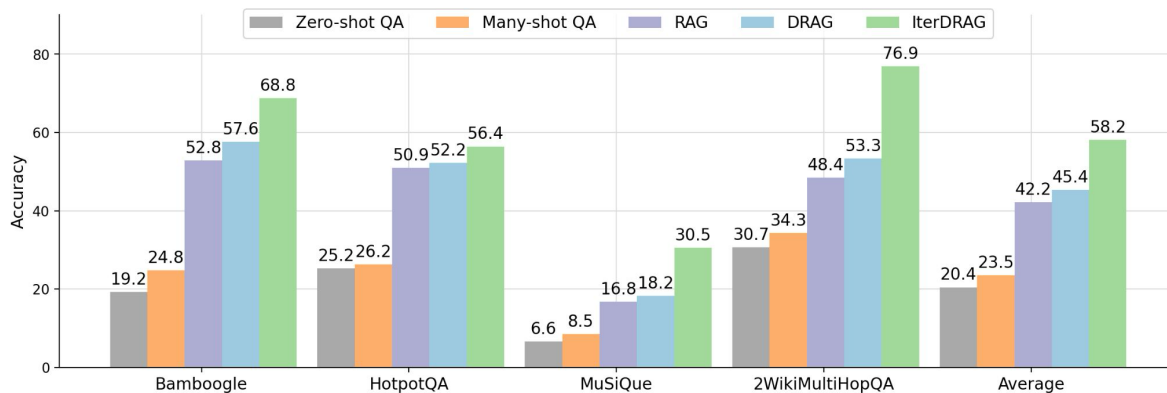
DRAG excels with shorter effective context lengths but IterDRAG scales more effectively for longer ones

RAG Performance and Inference Computation Scale

Comparing Strategies

Comparing the optimal performance with effective context length L_{max} up to 5M

DRAG and IterDRAG can achieve better performance than baselines



DRAG and IterDRAG
with the optimal
configuration leverage
the context window
better than RAG.

DRAG and IterDRAG
with the optimal
configuration leverage
the context window
better than RAG.

*How to find the optimal
configuration without
brute-force?*

Inference Computation Allocation for RAG

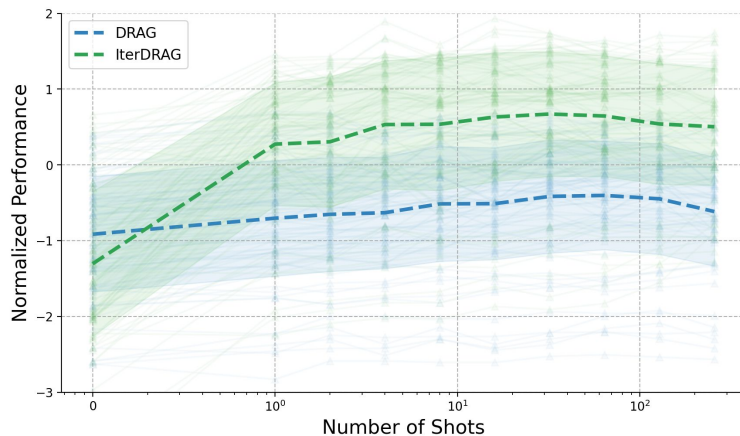
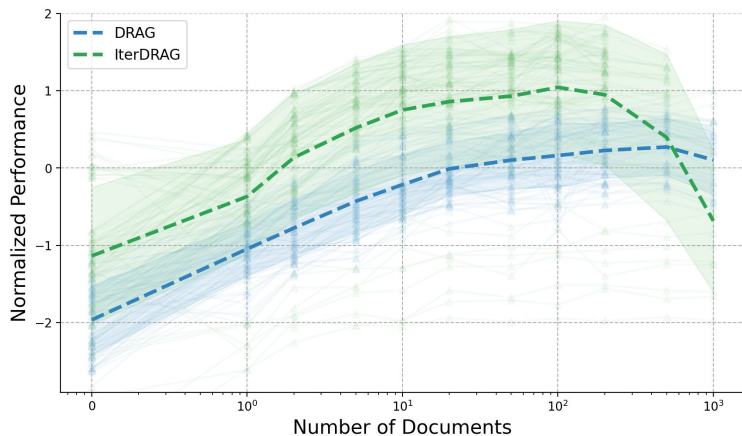
4

RAG Performance and Inference Computation Scale

RAG Performance vs. Different Parameters

Plotting DRAG and IterDRAG performance vs. individual parameter: number of documents and number of demonstrations.

1. *Number of documents* is more helpful than *number of demonstrations*, as the curve has steeper slope
2. IterDRAG benefits more from increasing number of demonstration



Inference Computation Allocation for RAG

Introducing a Quantitative Model

Denote the parameters of the strategies as $\theta = [k, m, n]^T$ we can formulate the computation allocation model as

$$\sigma^{-1}(P(\theta)) \approx (a + b \odot i)^T \log(\theta) + c$$

where

- a, b, c are parameters to learn;
- i is a vector of informativeness that can be easily estimated for each dataset individually;
- σ is a link function and $\odot$ refers to element-wise product.

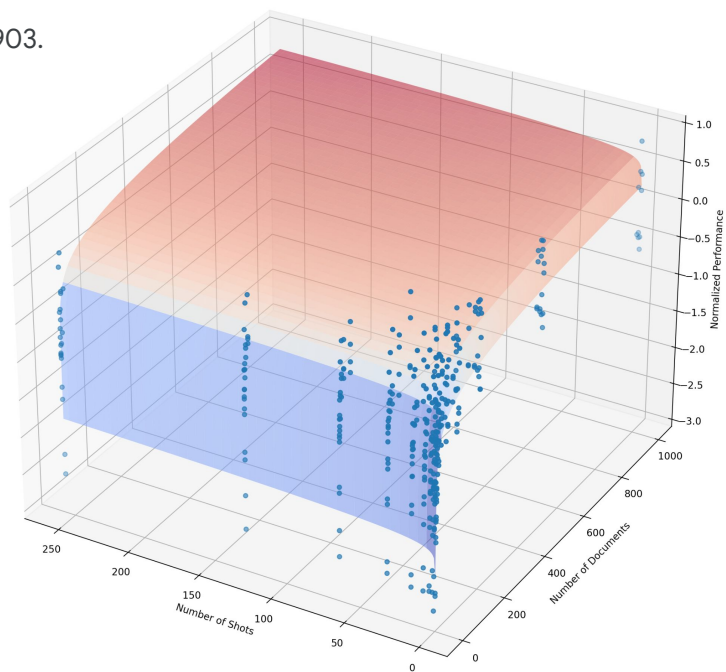
The informativeness i include informativeness for adding a document, and informativeness for adding a demonstration, estimated by

1. i_{doc} = performance difference between $k=1$ document and $k=0$ document
2. i_{shot} = performance difference between $m=1$ demo and $m=0$ demo respectively
3. i_{iter} = 0 as we do not find an accurate way to estimate informativeness of adding an iteration

Inference Computation Allocation for RAG

Estimated Model

The estimated model has an R^2 of 0.903.

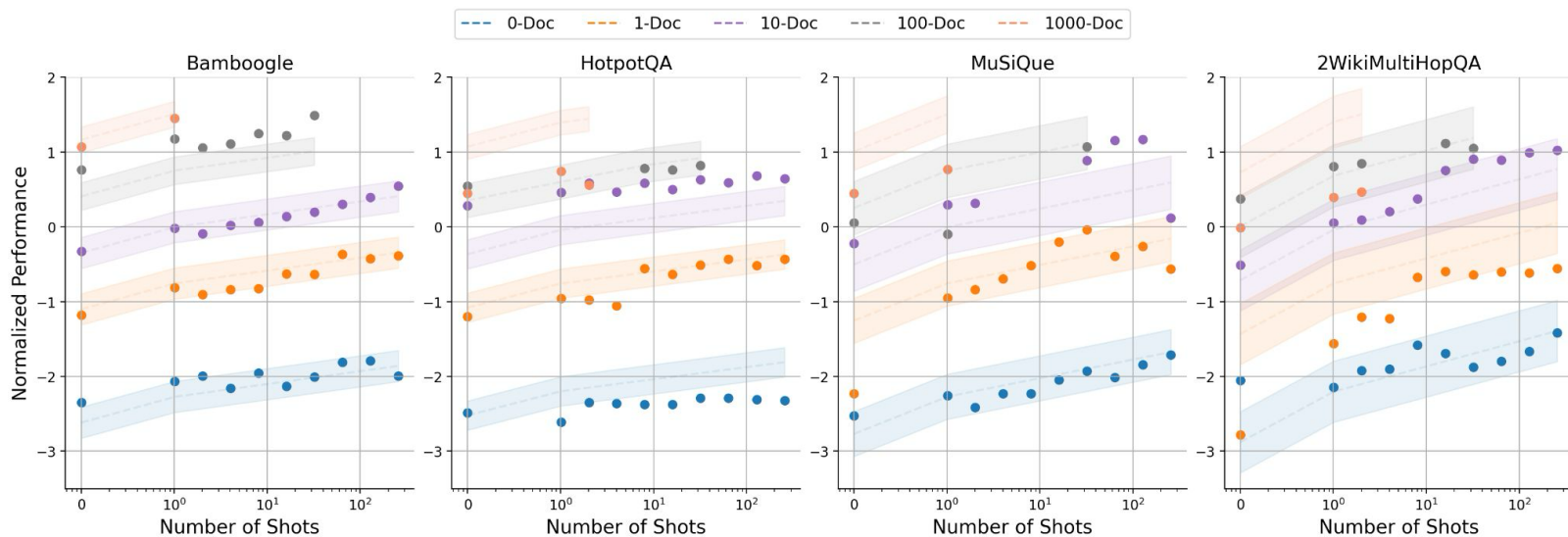


Inference Computation Allocation for RAG

Estimated Model

The estimated model has an R^2 of 0.903.

Plot the model estimation (shaded area) vs. the actual performance for a few slices of configurations:



Inference Computation Allocation for RAG

Predict the Optimal Configuration

Evaluate how well the model generalize in two different ways:

1. Use 3 datasets to fit the model and predict the optimal on the other one
2. Use data from shorter effective context lengths to fit the model and predict the optimal on longer effective context lengths

*Baseline = always 8-shot and fill the context length with as many documents as possible

Generalization to other datasets

	Bamboogle			HotpotQA			MuSiQue			2WikiMultiHopQA		
	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc
Baseline	49.6	58.8	51.2	46.3	60.2	51.4	14.9	24.7	16.9	46.5	53.7	51.6
Predict	64.0	75.6	68.0	47.8	63.3	55.3	19.3	32.5	29.3	60.8	72.4	74.9
Oracle	65.6	75.6	68.8	48.7	63.3	55.3	22.2	34.3	30.5	65.7	75.2	76.4

Generalization to longer context lengths

	16k → 32k			32k → 128k			128k → 1M			1M → 5M		
	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc	EM	F1	Acc
Baseline	37.4	47.6	40.4	39.0	49.5	42.2	39.3	49.3	42.8	44.5	55.4	49.8
Predict	37.4	48.2	41.0	41.2	52.0	45.4	48.0	60.9	56.9	47.9	59.8	55.2
Oracle	39.2	49.8	42.7	46.9	59.0	55.1	50.5	62.1	57.7	51.7	62.6	58.1

Summary

5

Summary

01

We comprehensively investigate inference scaling for RAG in the regime of long-context LLMs.

We use two inferences scaling strategies: DRAG and IterDRAG.

02

With an enriched toolbox to scale inference computation, we observe that the optimal RAG performance can scale almost linearly with the order of magnitude of inference computation.

This is different from previous observations where the RAG performance tends to saturate or drop when only scaling the number of documents.

03

We develop a quantitative model to predict RAG performance for a specific inference parameter configuration.

The model can be used to *find the optimal configuration* for a given inference computation budget.

Experiments show reasonable predictive power.



Thank you.